



A derivative-free optimization algorithm for the efficient minimization of functions obtained via statistical averaging

Pooriya Beyhaghi¹ · Ryan Alimo¹ · Thomas Bewley¹

Received: 29 September 2016 / Published online: 4 February 2020
© Springer Science+Business Media, LLC, part of Springer Nature 2020

Abstract

This paper considers the efficient minimization of the infinite time average of a stationary ergodic process in the space of a handful of design parameters which affect it. Problems of this class, derived from physical or numerical experiments which are sometimes expensive to perform, are ubiquitous in engineering applications. In such problems, any given function evaluation, determined with finite sampling, is associated with a quantifiable amount of uncertainty, which may be reduced via additional sampling. The present paper proposes a new optimization algorithm to adjust the amount of sampling associated with each function evaluation, making function evaluations more accurate (and, thus, more expensive), as required, as convergence is approached. The work builds on our algorithm for Delaunay-based Derivative-free Optimization via Global Surrogates (Δ -DOGS, see JOGO <https://doi.org/10.1007/s10898-015-0384-2>). The new algorithm, dubbed α -DOGS, substantially reduces the overall cost of the optimization process for problems of this important class. Further, under certain well-defined conditions, rigorous proof of convergence to the global minimum of the problem considered is established.

Keywords Derivative-free optimization · Statistical averaging · Delaunay triangulation

✉ Pooriya Beyhaghi
p.beyhaghi@gmail.com

Ryan Alimo
shahrouz.ryan.alimo@gmail.com

Thomas Bewley
bewley@eng.ucsd.edu

¹ Flow Control Lab, University of California, San Diego, San Diego, USA

1 Introduction

In this paper, the Delaunay-based derivative-free optimization algorithm developed in [1, 2, 12–14, 47], dubbed Δ -DOGS, is modified to minimize an objective function, $f(x) : \mathbb{R}^n \rightarrow \mathbb{R}$, which cannot be evaluated directly, but can be approximated to a tuneable degree of precision. As the precision of any function evaluation is increased, the computational (or experimental) cost of that function evaluation is increased accordingly. An example for this type of objective function is the infinite-time average of a discrete-time ergodic process $g(x, k)$ for $k = 1, 2, 3, \dots$ such that

$$f(x) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{k=1}^N g(x, k), \quad (1a)$$

where, for $k \geq \bar{k}$, $g(x, k)$ is assumed to be statistically stationary. The feasible domain in which the optimal design parameter vector $x \in \mathbb{R}^n$ is sought is a bound constrained domain

$$L = \{x | a \leq x \leq b\} \quad \text{where } a < b \in \mathbb{R}^n. \quad (1b)$$

In practice, the precise numerical determination of $f(x)$ for any x is not possible, as this would require infinite time averaging; $f(x)$ can only be approximated as the average of $g(x, k)$ over some *finite* number of samples N . The truth function $f(x)$ is typically a smooth function of x , though it is often nonconvex; computable approximations of $f(x)$, however, are generally *nonsmooth* in x , as the truncation error (associated with the fact any approximation of $f(x)$ must be computed with finite N) is effectively decorrelated from one approximation of $f(x)$ to the next.

Minimizing (1a) within the feasible domain (1b) is the subject of interest in host of practical applications, such as the optimization of stiffness and shape parameters [25], feedback control gains in mechanical systems [10] and manufacturing processes involving turbulent flows [37], etc.

One interesting class of global optimization algorithms is the Direct (DIviding RECTangles) method which was developed in [22] for optimizing Lipschitz functions. These methods are extended in [38] to the case of any possible semi-metric by simultaneously considering the subspaces that can contain the optimum. The methods are deterministic optimization algorithms which was designed for problems with exact function evaluation. Later, these methods were modified to solve for objective functions obtained from noisy measurements in [19, 23, 45].

A second relevant class of global optimization algorithms are branch and bound methods [28], which partition the search domain, then characterize the most promising partition in which to perform the next function evaluation. These methods were initially designed for optimizing objective functions in which exact function evaluations were possible; however, the methods in [33] modified it to address problems with noisy function evaluations.

Another class of methods are polynomial optimization algorithms [24] which globally solves an optimization algorithm using Lasserre type Moment-SOS relaxations. The methods in [27] extends these algorithms to address stochastic

optimization problems. However, these methods are limited to the problems where $f(x)$ is a polynomial function of x .

Derivative-free optimization methods is discussed in [5, 17, 18, 35], and indeed appear to be the most promising class of approaches for problems of the present form. These methods are implemented for shape optimization in airfoil design [25], as well as in online optimization [23]. With such methods, only values of the function evaluations themselves are used, and neither a derivative nor its estimate is needed. The best methods of this class strive to keep function evaluations far apart in parameter space until convergence is approached, thereby mitigating somewhat the effect of uncertainty in the function evaluations. This class of methods generally handles bound constraints quite well, and may be used to globally minimize the function of interest. Moreover, some advance algorithms [3, 4, 6] in this class can handle problems with nonlinear constraints. However, this class of method scales poorly with the dimension of the problem. The surrogate management framework [15, 35, 42] and Bayesian algorithms [30–32, 36, 39] are amongst the best derivative-free methods available today, and are implemented for minimizing a problem of the form in (1) in [25, 26, 40, 43]. In this class, the method developed in [39] develops a promising Bayesian approach, in a manner which increases the sampling of new measurements as convergence is approached. However, this method does not selectively refine existing measurements, which is a key contributor to the efficiency of the algorithm developed herein.

In this paper, a provably globally convergent (under the appropriate assumptions) new optimization approach is developed for problems of the form given in (1). The structure of the remainder of the paper is as follows: Sect. 2 briefly reviews the key features of the Δ -DOGS(Z) algorithm developed in [12], upon which the present paper is built. Section 3 lays out all of the new elements that compose the new optimization approach, as well as the new algorithm itself, dubbed α -DOGS. Section 4 analyzes the convergence properties of the new algorithm, and establishes conditions which are sufficient to guaranty its convergence to the global minimum. Section 5 applies the new algorithm to a selection of model problems in order to illustrate its behavior. Some conclusions are presented in Sect. 7.

2 Delaunay-based optimization coordinated with a grid: Δ -DOGS(Z)

This section presents a simplified version of the Δ -DOGS(Z) algorithm, the full version of which is given as Algorithm 2 of [12], where it is analyzed in detail. The Δ -DOGS(Z) algorithm is a grid-based acceleration of the Δ -DOGS algorithm originally developed in [14], and is designed to minimize problems in which precise function evaluations are available, while avoiding an accumulation of unnecessary function evaluations on the boundary of the feasible domain.

The optimization problem considered in this section is the minimization of an objective function $f(x)$, approximations of which are assumed to be available, in the feasible domain $L = \{x | x \in \mathbb{R}^n, a \leq x \leq b\}$. At each iteration of the simplified Δ -DOGS(Z) algorithm considered here, a metric based on an interpolation of existing function evaluations, and a model $e(x)$ of the “remoteness” of

any point $x \in L$ from the available datapoints at that iteration (which in the Δ -DOGS(Z) algorithm characterizes the uncertainty of the interpolation) is minimized to obtain the location of the next point $x \in L$ at which the function will be evaluated. The interpolation and remoteness functions at iteration k are denoted here by $p^k(x)$ and $e^k(x)$, the latter of which is defined below. Note that the function $e^k(x)$ is called an “uncertainty” function in [12–14]; we use the name “remoteness” function for the same construction in the present paper, as it plays a slightly different role in the sections that follow.

Definition 1 Consider S as a set of feasible points which includes the vertices of L , and Δ as a Delaunay triangulation of S . Then, for each simplex $\Delta_i \in \Delta$, the local remoteness function is defined as:

$$e_i(x) = R_i^2 - \|x - Z_i\|^2, \quad (2a)$$

where R_i and Z_i are the circumradius and circumcenter of Δ_i . The global remoteness function $e(x)$ is a piecewise quadratic function which is nonnegative everywhere and goes to zero at the datapoints, and is defined as follows:

$$e(x) = e_i(x) \quad \forall x \in \Delta_i. \quad (2b)$$

The remoteness function $e(x)$ has a number of properties which are established in Lemmas [2:5] in [14], as listed below.

- (a) The remoteness function $e(x) \geq 0$ for all $x \in L$, and $e(x) = 0$ for all $x \in S$.
- (b) The remoteness function $e(x)$ is continuous and Lipschitz.
- (c) The remoteness function $e(x)$ is piecewise quadratic with Hessian of $-2I$.
- (d) The remoteness function $e(x)$ is equal to the maximum of the local remoteness functions:

$$e(x) = \max_{1 \leq i \leq N_s} e_i(x), \quad (3)$$

where N_s is the number of simplices.

We now review the definition of the Cartesian grid over the feasible domain, as discussed further in [12].

Definition 2 Taking $N_\ell = 2^\ell$, the Cartesian grid of level ℓ over the feasible domain $L = \{x | a \leq x \leq b\}$, denoted L_ℓ , is defined as follows:

$$L_\ell = \left\{ x | x_j = a_j + \frac{b_j - a_j}{N_\ell} \cdot z, \quad z \in \{1, 2, \dots, N_\ell\}, \quad j \in \{1, 2, \dots, n\} \right\}.$$

The *quantization* of a point x onto the grid L_ℓ , denoted x_q^ℓ , is a point on the grid L_ℓ which has the minimum distance from x . Note that this quantization process might have multiple solutions; any of these solutions is acceptable. The maximum quantization error of the grid, δ_{L_ℓ} , is defined as follows:

$$\delta_{L_\ell} = \max_{x \in L_\ell} \|x - x_q^\ell\| = \frac{\|b - a\|}{2N_\ell}. \quad (4)$$

Remark 1 Three important properties of the Cartesian grid for the present purposes follow.

- (a) The grid of level ℓ covering the feasible domain L in an n dimensional space has $(N_\ell + 1)^n$ grid points.
- (b) $\lim_{\ell \rightarrow \infty} \delta_{L_\ell} = 0$.
- (c) If x_q^ℓ is a quantization of x onto L_ℓ , then $A_a(x) \subseteq A_a(x_q^\ell)$, where $A_a(x)$ is the set of active bound constraints at x .

Algorithm 1 presents a strawman form of the Δ -DOGS(Z) algorithm. A significant refinement of this algorithm is presented as Algorithm 2 of [12], together with its proof of convergence and its implementation on model problems. The key refinement in [12] of the strawman algorithm presented here is the identification of two different sets of points in L , dubbed S_E (on which function evaluations are performed) and S_U (on which function evaluations are not yet performed); the latter set proves useful in [12] to regularize the triangulation. The algorithm proposed in Sect. 3 below effectively generalizes this notion, of evaluation points at which the function has been evaluated, and support points at which the function has not yet been evaluated, to the quantification and control over the *extent* of sampling performed for any given approximation of $f(x)$.

Algorithm 1 Strawman form of the grid-accelerated Delaunay-Based optimization algorithm, Δ -DOGS(Z), presented as Algorithm 2 in [12], which assumes that precise function evaluations are available.

- 1: Set $k = 0$ and initialize ℓ . Take the set of initialization points S^0 as all 2^n vertices of the feasible domain L , together with any user-supplied points of interest (quantized onto the grid L_0), and perform function evaluations at all points in S^0 .
 - 2: Calculate (or, for $k > 0$, update) an appropriate interpolating function $p^k(x)$ through all points in S^k .
 - 3: Calculate (or, for $k > 0$, update) a Delaunay triangulation Δ^k over all of the points in S^k .
 - 4: Find z as a global minimizer of $s_c^k(x) = p^k(x) - K e^k(x)$ in L , and take z_ℓ as its quantization onto the grid L_ℓ .
 - 5: If $z_\ell \notin S^k$, then take $S^{k+1} = S^k \cup \{z_\ell\}$ calculate $f(z_\ell)$, increment k , and repeat from 2.
 - 6: Otherwise, take $S^{k+1} = S^k$, increment both k and ℓ , and repeat from 2.
-

3 Delaunay-based optimization of a time-averaged value: α -DOGS

This section presents the essential elements of the new optimization algorithm, dubbed α -DOGS, which is designed to efficiently minimize a function $f(x)$ given by (1a) within the feasible domain L defined by (1b). We begin by introducing some fundamental concepts.

Definition 3 Take S as a finite set of points x_i , for $i = 1, \dots, M$, at which the function $f(x)$ in (1a) has been approximated; drawing a parallel to the nomenclature commonly used in estimation theory, we refer to any such approximation of $f(x)$, developed with a finite number of samples N_i , as a *measurement*, denoted y_i :

$$y_i = y(x_i, N_i) = \frac{1}{N_i} \sum_{k=1}^{N_i} g(x_i, k). \quad (5)$$

Any such measurement y_i has a finite uncertainty associated with it, which can be reduced by increasing N_i .

Remark 2 For many problems, there is an initial transient such that, for $k < \bar{k}$, the assumption of stationarity of $g(x, t_k)$ is not valid. For such problems, the initial transient in the data can be detected using the approach developed in [7] and set aside, and the signal considered as stationary thereafter. In such problems, to increase the speed of convergence of the statistics, the finite sum used for averaging the samples in (5) is modified to begin at $\bar{k}$ instead of beginning at 1.

Since, for $k > \bar{k}$, $g(x, k)$ is statistically stationary, each measurement y_i is an *unbiased estimate* of the corresponding value of $f(x_i)$. We assume that a model for the *standard deviation* quantifying the uncertainty of this measurement, denoted $\sigma_i = \sigma(x_i, N_i)$, is also available. Since $g(x, t)$ is a stationary ergodic process, for any point $x_i \in L$,

$$\lim_{N_i \rightarrow \infty} y(x_i, N_i) = f(x_i), \quad \lim_{N_i \rightarrow \infty} \sigma(x_i, N_i) = 0. \quad (6)$$

Remark 3 If a stationary ergodic process $g(x, k)$ at some point $x_i \in L$ is *independent and identically distributed* (IID), then $\sigma(x_i, N_i) = \sigma(x_i, 1)/\sqrt{N_i}$; otherwise, estimates of $\sigma(x_i, N_i)$ can be developed using standard uncertainty quantification (UQ) procedures, such as those developed in [8, 9, 11, 29, 34, 44]. The discrete-time process $g(x, k)$ may often be obtained by sampling a continuous-time process $g(x, t)$ at timesteps $t_k = kh$ for some appropriate sample interval h . For sufficiently large h , the samples of this continuous-time process $g(x, k)$ are often essentially IID; however, with the appropriate UQ procedures in place, significantly smaller sample intervals h will lead to a given degree of convergence in a substantially shorter period of time t , albeit with increased storage.

Definition 4 Define S as a set of measurements y_i , for $i = 1, 2, \dots, M$, at corresponding points x_i and with standard deviation σ_i . We will call a regression $p(x)$ for this set of measurements a *strict regression* if, for some constant β ,

$$|p(x_i) - y_i| \leq \beta \sigma_i, \quad \forall 1 \leq i \leq M. \quad (7)$$

Based on the concepts defined above, Algorithm 2 presents our algorithm to efficiently and globally minimize a function of the form (1a) within a feasible domain defined by (1b). At each iteration k of this algorithm, S^k denotes the set of

M points x_i , for $i = 1, 2, \dots, M$, at which measurements have so far been made; for each point $x_i \in S^k$, $y_i = y(x_i, N_i)$ denotes the measured value, $\sigma_i = \sigma(x_i, N_i)$ denotes the uncertainty of this estimate, and N_i quantifies the sampling performed thus far at point x_i . Note that the values of M , N_i , y_i , ℓ , α , and K are all updated from time to time as the iterations proceed, and are thus annotated with a k superscript at various points in the analysis of Sect. 4 to remove ambiguity. Akin to Algorithm 1, at iteration k , $p^k(x)$ is assumed to be a strict regression (for some value of β) of the current set of measurements y_i , and ℓ is the current grid level.

Algorithm 2 The new optimization algorithm, dubbed α -DOGS, for minimizing the function $f(x)$ in (1a) within the feasible domain L defined in (1b).

- 1: Set $k = 0$ and initialize the algorithm parameters α , K , γ , β , ℓ , N^0 , and N^δ as discussed in §3. Take the initial set of M sampled points, S^0 , as the 2^n vertices of the feasible domain $L = \{x | a \leq x \leq b\}$ together with any user-supplied points of interest quantized onto the grid L_ℓ . Take $N_i = N^0$ for $i = 1, \dots, M$, and compute an initial measurement $y_i = y(x_i, N^0)$ and corresponding uncertainty $\sigma_i = \sigma(x_i, N^0)$ for each point $x_i \in S^0$.
- 2: Calculate a strict regression $p^k(x)$ for all M available measurements.
- 3: Calculate (or, for $k > 0$, update) a Delaunay triangulation Δ^k over all of the points in S^k .
- 4: Determine x_j as the minimizer (and, j as the corresponding index), over all $x_i \in S^k$, of the discrete search function $s_d^k(x_i)$, defined as follows:

$$s_d^k(x_i) = \min\{p^k(x_i), 2y_i - p^k(x_i)\} - \alpha^k \sigma_i^k \quad \text{for } i = 1, \dots, M. \quad (8)$$

- 5: Noting the definition of $e^k(x)$ in (2), determine z as the minimizer, over all $x \in L$, of the continuous search function $s_c^k(x)$, defined as follows:

$$s_c^k(x) = p^k(x) - K^k e^k(x) \quad \text{for } x \in L. \quad (9)$$

Denote z_ℓ as the quantization of z onto the grid L_ℓ .

- 6: If $s_c^k(z) > s_d^k(x_j)$ and $N_j < \gamma 2^\ell$ where N_j is the current number of samples taken at x_j , then take N^δ additional samples at x_j , update $N_j \leftarrow N_j + N^\delta$, update the measurement $y_j = y(x_j, N_j)$ and uncertainty $\sigma_j = \sigma(x_j, N_j)$, increment k , and repeat from 2.
 - 7: Otherwise, if $z_\ell \notin S^k$, then set $x_{M+1} = z_\ell$ and $S^{k+1} = S^k \cup \{x_{M+1}\}$, take $N_{M+1} = N^0$, compute the measurement $y_{M+1} = y(x_{M+1}, N^0)$ and uncertainty $\sigma_{M+1} = \sigma(x_{M+1}, N^0)$, increment M and k , and repeat from 2.
 - 8: Otherwise (i.e., if $z_\ell \in S^k$), adjust the algorithm parameters such that $\alpha^{k+1} \leftarrow \alpha^k + \alpha^\delta$ and $K^{k+1} \leftarrow 2K^k$, increment both ℓ and k , and repeat from 2.
-

At each iteration of Algorithm 2, there are three possible situations, corresponding to three of the numbered iterations of this algorithm:

- (6) The sampling of an existing measurement is increased. This is called a *supplemental sampling* iteration.
- (7) A new point is identified, and an initial measurement at this point is added to the dataset. This is called an *identifying sampling* iteration.
- (8) The mesh coordinating the problem is refined and the algorithm parameters α and K adjusted. This is called a *grid refinement* iteration.

Figure 1 illustrates supplemental sampling and identifying sampling iterations of Algorithm 2.

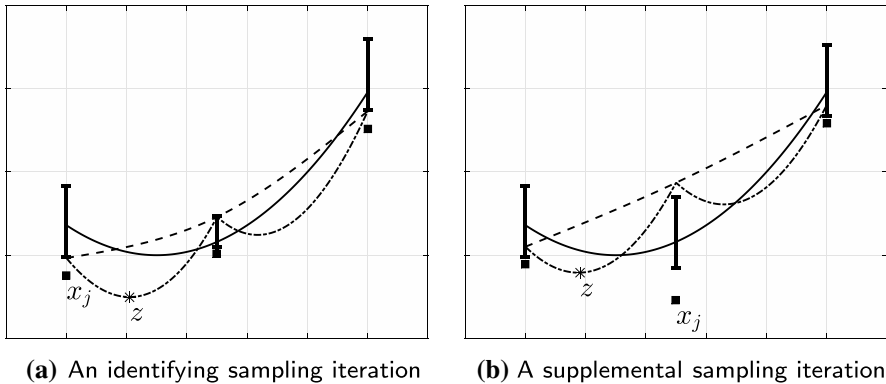


Fig. 1 Representation of one iteration in Algorithm 2 in different situations: (solid line) truth function $f(x)$, (dashed line) the regression $p^k(x)$, (dash-dot line) continuous search function $s_c^k(x)$, (closed squares) $s_d^k(x)$, and (asterix) z . In **(a)**, $s_c^k(z) < s_d^k(x_j)$; it is thus an identifying sampling iteration. In **(b)**, $s_c^k(z) > s_d^k(x_j)$; it is thus an supplemental sampling iteration. Horizontal axis is the x coordinate and the vertical axis is the function value $f(x)$

Algorithm 2 depends upon a handful of algorithm parameters, the selection of which affects its rate of convergence, explored in Sect. 5, though not its proof of convergence, established in Sect. 4. The remainder of this section discusses heuristic strategies to tune these algorithm parameters, noting that this tuning is an application-specific problem, and alternative strategies (based on experiment or intuition) might lead to more rapid convergence for certain problems.

The first task encountered during the setup of the optimization problem is the definition of the design parameters. Note that the feasible domain considered during the optimization process is characterized by simple upper and lower bounds for each design parameter; normalizing all design parameters to lie between 0 and 1 is often helpful.

The second challenge is to scale the function $f(x)$ itself, such that the range of the normalized function $f(x)$ over the feasible domain L is about unity. If an estimate of the actual range of $f(x)$ is not available a priori, we may estimate it at any given iteration using the available measurements. Following this approach, at any iteration k with available measurements $\{y_1, y_2, \dots, y_M\}$, all measurements y_i , as well as the corresponding uncertainty of these measurements σ_i , may be scaled by a factor r_s wherever used in iteration k of Algorithm 2 where, for that iteration, r_s is computed such that

$$r = \frac{1}{\max_{1 \leq i \leq M} \{y_i\} - \min_{1 \leq i \leq M} \{y_i\}}, \quad r_s = r_l + R(r - r_l) - R(r - r_u),$$

where $R(x)$ is the ramp function. So defined, r_s is a “saturated” version of the factor r , constrained to lie in the range $r_l \leq r_s \leq r_u$. Note that scaling the y_i and σ_i does not interfere with the proof of the convergence of the algorithm, provided in Sect. 4, but

can improve its performance. In the numerical simulations performed in Sect. 5, we take $r_l = 10^{-3}$ and $r_u = 10^3$.

For problems which are IID with no initial transient, $N^0 = N^\delta = 1$ is a reasonable starting point; increasing N^0 and N^δ ultimately reduces the number of iterations of the algorithm (and, thus, the number of Delaunay triangulations) required for convergence, but generally increases slightly the total amount of sampling performed. Suggested values of other algorithm parameters, which work well in the numerical simulations reported in Sect. 5 but the values of which do not affect the proof of convergence provided in Sect. 4, include $\alpha^0 = \alpha^\delta = 0.5$, $K^0 = 0.5$, $\ell^0 = 3$, $\beta = 4$, and $\gamma = 100$.

4 Analysis of α -DOGS

We now analyze the convergence of Algorithm 2. We first present some preliminary definitions.

Definition 5 The point $\eta^k \in S$ is called the *candidate* point at iteration k if

$$\eta^k \in \operatorname{argmin}_{z \in S^k} \{y_k(z) + \alpha^k \sigma_k(z)\}. \quad (10)$$

Define $f(x^*)$ as the global minimum of $f(x)$ in L , then the *regret* is defined as

$$r_k = f(v_k) - f(x^*). \quad (11)$$

The definition of the regret given above is common in the optimization literature (see, e.g., [16, 23, 40]). We show in this section that, under the following assumptions, the regret of the optimization process governed by Algorithm 2 will converge to zero:

Assumption 1 A constant $\hat{K}$ exist such that, for all $k > 0$ and $x \in L$,

$$-2\hat{K}I \leq \nabla^2 p^k(x) \leq 2\hat{K}I, \quad -2\hat{K}I \leq \nabla^2 f(x) \leq 2\hat{K}I,$$

where I is the identity matrix.

Assumption 2 There is a real continuous and monotonically increasing function $E : \mathbb{R}^+ \rightarrow [0, Q]$, which has the following properties:

- $E(0) = 0$ and $\lim_{r \rightarrow 0^+} E(r) = 0$ and $\lim_{r \rightarrow \infty} E(r) = Q$.
- For all $x \in L$ and $N \in \mathbb{N}$, we have:

$$\left| \frac{1}{N} \sum_{k=1}^N g(x, k) - f(x) \right| = |y(x, N) - f(x)| \leq E(\sigma(x, N)). \quad (12)$$

Assumption 3 There are real numbers $\alpha > 0$ and $\theta \in (0, 1]$, such that

$$\sigma(x, N) \leq \alpha N^{-\theta} \quad \forall x \in \mathbb{R}, N \in \mathbb{N}. \quad (13)$$

Note that, if the stationary process $g(x, k)$ has a short-range dependence, like ARMA processes, the parameter $\theta = 0.5$. However, for different θ , this general model can also handle stationary processes with long-range dependence, like Fractional ARMA processes. Moreover, at any given point $x \in L$, $\sigma(x, N)$ is a monotonically nonincreasing function of N . This condition means that, by increasing the averaging interval, the uncertainty of the estimate must not increase.

Remark 4 Assumption 2 is a stronger condition than ergodicity of $g(x, k)$. Recall that ergodicity of $g(x, k)$ is equivalent to the convergence of $y(x, N)$ to $f(x)$ as $N \rightarrow \infty$, which is the straightforward outcome of Assumptions 2 and 3. In Assumption 2, the convergence of the sample mean is assumed to be bounded by a function of $\sigma(x, N)$ of the form specified.

Lemma 1 For any point $x \in L$ and real positive number $0 < \varepsilon < Q$,

$$|y(x, N) - f(x)| - \frac{Q}{E^{-1}(\varepsilon)} \sigma(x, N) \leq \varepsilon, \quad (14)$$

Proof If $\sigma(x, N) \leq E^{-1}(\varepsilon)$ then, since $E(x)$ is an increasing function, by (12), we have:

$$E(\sigma(x, N)) \leq E(E^{-1}(\varepsilon)) = \varepsilon \quad \Rightarrow \quad |y(x, N) - f(x)| \leq \varepsilon.$$

Otherwise, $\sigma(x, N) > E^{-1}(\varepsilon)$; thus, again by (12), we have

$$\frac{Q}{E^{-1}(\varepsilon)} \sigma(x, N) \geq Q \quad \Rightarrow \quad |y(x, N) - f(x)| - Q \leq E(\sigma_N) - Q \leq 0,$$

Thus, (14) is verified for both cases. $\square$

Lemma 2 During the execution of Algorithm 2, there are an infinite number of mesh refinement iterations.

Proof This lemma is shown by contradiction. If Algorithm 2 has a finite number of mesh refinement iterations, then there is an integer number $\bar{\ell}$ such that the mesh $L_{\bar{\ell}}$ contains all datapoints obtained by the algorithm. Since the number of datapoints on this mesh is finite, only a finite number of points must be considered, which leads to having a finite number of identifying sampling iterations.

Since the number of identifying sampling and mesh refinement iterations are finite, there must be an infinite number of supplemental sampling iterations. At each supplemental sampling iteration, the averaging length of the estimate at an existing datapoint is incremented by $N^\delta \geq 1$. Since only a finite number of points is considered, a datapoint exists for which the estimate is improved for an infinite number of supplemental sampling iterations. As a result, there is an supplemental sampling

iteration, such that $N_j > \gamma 2^{\bar{\epsilon}}$, which is in contradiction with the assumption of having finite number of mesh refinement iterations. $\square$

Note that the following short Lemma and proof, which are necessary for this development, are copied directly from [12].

Lemma 3 Consider $G(x)$ as a twice differentiable function such that $\nabla^2 G(x) - 2K_1 I \preceq 0$, and $x^* \in L$ as a local minimizer of $G(x)$ in L . Then, for each $x \in L$ such that $A_a(x^*) \subseteq A_a(x)$, we have:

$$G(x) - G(x^*) \leq K_1 \|x - x^*\|^2. \quad (15)$$

Proof Define function $G_1(x) = G(x) - K_1 \|x - x^*\|^2$. By construction, $G_1(x)$ is concave; therefore,

$$\begin{aligned} G_1(x) &\leq G_1(x^*) + \nabla G_1(x^*)^T (x - x^*), \\ G_1(x^*) &= G(x^*), \quad \nabla G_1(x^*) = \nabla G(x^*), \\ G(x) &\leq G(x^*) + \nabla G(x^*)^T (x - x^*) + K_1 \|x - x^*\|^2. \end{aligned}$$

Since the feasible domain is a bounded domain, the constrained qualification holds; therefore, x^* is a KKT point. Therefore, using $A_a(x^*) \subseteq A_a(x)$ leads to $\nabla G(x^*)^T (x - x^*) = 0$, which verifies (15). $\square$

Lemma 4 Consider z , x_j , and x^* as global minimizers of $s_c^k(x)$, $s_d^k(x)$, and $f(x)$, respectively. Note that $s_d^k(x)$ is only defined for the points in S^k , but $s_c^k(x)$ and $f(x)$ are defined over the feasible domain L . Define M_k as:

$$M_k = \min\{s_c^k(z) - f(x^*), s_d^k(x_j) - f(x^*)\}. \quad (16)$$

Then,

$$\limsup_{k \rightarrow \infty} M_k \leq 0. \quad (17)$$

Proof By Lemma 2, there are infinite number of mesh refinement iterations during the execution of Algorithm 2. Thus,

$$\lim_{k \rightarrow \infty} K^k = \infty, \quad \lim_{k \rightarrow \infty} \alpha^k = \infty. \quad (18)$$

As a result, for any $0 < \epsilon < Q$, there is a k_ϵ such that, if $k > k_\epsilon$, then

$$K^k \geq 3\hat{K} \quad \text{and} \quad \alpha^k \geq \frac{2Q}{E^{-1}(\epsilon)}. \quad (19)$$

Consider $\Delta_{x^*}^k$ as a simplex in Δ^k , a Delaunay triangulation for S^k , that contains x^* . Define $M(x) : \Delta_{x^*}^k \rightarrow \mathbb{R}$ as the unique linear function in $\Delta_{x^*}^k$ such that

$$M(V_j^k) = 2f(V_j^k) - p^k(V_j^k),$$

where V_j^k are the vertices of $\Delta_{x^*}^k$. Define $G(x) : \Delta_{x^*}^k \rightarrow \mathbb{R}$ as follows:

$$G(x) = s_c^k(x) + M(x) - 2f(x) = p^k(x) + M(x) - 2f(x) - K^k e^k(x).$$

By construction, $G(V_j^k) = 0$. Moreover:

$$\nabla^2 G(x) = \nabla^2 \{p^k(x) - 2f(x)\} + 2K^k I.$$

Using Assumption 1 and (19), $G(x)$ is strictly convex in simplex $\Delta_{x^*}^k$. Since $G(x) = 0$ at the vertices of $\Delta_{x^*}^k$, then $G(x^*) \leq 0$. Moreover, since $M(x)$ is a linear function, then

$$\begin{aligned} \min_{x \in S^k} [2f(x) - p^k(x)] &\leq \min_{1 \leq j \leq n+1} [2f(V_j^k) - p^k(V_j^k)] \leq M(x^*), \\ s_d^k(x) &\leq [2f(x) - p^k(x)] + 2(y_k(x) - f(x)) - \alpha^k \sigma(x). \end{aligned} \quad (20)$$

Using (14) in Lemma 1 and (19) leads to:

$$2y_k(x) - 2f(x) - \alpha^k \sigma(x) \leq 2\varepsilon. \quad (21)$$

Combining (20) and (21) leads to:

$$s_d^k(x) \leq [2f(x) - p^k(x)] + 2\varepsilon.$$

Since x_j is the minimizer of the $s_d^k(x)$,

$$s_d^k(x_j) \leq \min_{x \in S^k} [2f(x) - p^k(x)] + 2\varepsilon \leq M(x^*) + 2\varepsilon.$$

Furthermore, z is the global minimizer of $s_c^k(x)$ and $G(x^*) \leq 0$; therefore,

$$\begin{aligned} s_c^k(z) &\leq s_c^k(x^*) \leq 2f(x^*) - M(x^*), \\ s_c^k(z) + s_d^k(x_j) &\leq 2f(x^*) + 2\varepsilon. \end{aligned} \quad (22)$$

Thus, for any $\varepsilon > 0$ and $k > \hat{k}_\varepsilon$, (22) is satisfied; therefore, (17) is verified. $\square$

Lemma 5 *If $\{k_1, k_2, \dots\}$ are the mesh refinement iterations of Algorithm 2, then*

$$\limsup_{i \rightarrow \infty} \left\{ y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) - f(x^*) + \alpha^{k_i} \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \right\} \leq 0, \quad \text{and} \quad (23a)$$

$$\lim_{i \rightarrow \infty} \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) = 0, \quad (23b)$$

where η^{k_i} is the candidate point at iteration k_i and x^* is a global minimizer of $f(x)$ in L .

Proof Consider z as a global minimizer of $s_c^{k_i}(x)$ in L , and z_ρ as its quantization on L_ρ . Since iteration k_i is a mesh refinement, $z_\rho \in S^{k_i}$. Consider $\Delta_j^{k_i}$ as a simplex in the Delaunay triangulation Δ^{k_i} which contains z . By property (d) of the remoteness function $e^{k_i}(x)$,

$$\begin{aligned}
e^{k_i}(z_\ell) &\geq e_j^{k_i}(z_\ell), \quad e^{k_i}(z) = e_j^{k_i}(z), \\
s_c^{k_i}(z_\ell) - s_c^{k_i}(z) &= p^{k_i}(z_\ell) - p^{k_i}(z) + K^{k_i}(e^{k_i}(z) - e^{k_i}(z_\ell)), \\
s_c^{k_i}(z_\ell) - s_c^{k_i}(z) &\leq p^{k_i}(z_\ell) - p^{k_i}(z) + K^{k_i}(e_j^{k_i}(z) - e_j^{k_i}(z_\ell)).
\end{aligned}$$

By Property (c) of Remark 1 concerning the Cartesian grid quantizer, $A_a(z) \subseteq A_a(z_\ell)$. According to Assumption 1 and Property (c) of the remoteness function introduced in Definition 1, $\nabla^2 \{p^{k_i}(x) - K^{k_i}e_j^{k_i}(x)\} - \{\hat{K} + 2K^{k_i}\}I \leq 0$; thus, by Lemma 3 and the fact (see Lemma 5 in [14] for proof) that z globally minimizes $p^{k_i}(x) - K^{k_i}e_j^{k_i}(x)$,

$$s_c^{k_i}(z_\ell) - s_c^{k_i}(z) \leq \{\hat{K} + 2K^{k_i}\} \|z_\ell - z\|^2. \quad (24)$$

Define δ_{k_i} as the maximum quantization error at iteration k_i , then $\|z_\ell - z\| \leq \delta_{k_i}$. On the other hand, $z_\ell \in S^{k_i}$, which leads to $s_c^{k_i}(z_\ell) = p^{k_i}(z_\ell)$, and

$$p^{k_i}(z_\ell) \leq s_c^{k_i}(z) + \{\hat{K} + 2K^{k_i}\} \delta_{k_i}^2. \quad (25)$$

At each mesh refinement iteration of Algorithm 2, there are two possibilities. In the first case, $s_c^{k_i}(z) \leq s_d^{k_i}(x_j)$; since x_j is a minimizer of $s_d^{k_i}(x)$, then

$$\begin{aligned}
p^{k_i}(z_\ell) &\leq s_d^{k_i}(z_\ell) + \{\hat{K} + 2K^{k_i}\} \delta_{k_i}^2, \\
s_d^{k_i}(z_\ell) &\leq p^{k_i}(z_\ell) - \alpha^{k_i} \sigma(z_\ell, N_{z_\ell}^{k_i}), \\
\sigma(z_\ell, N_{z_\ell}^{k_i}) &\leq \frac{\{\hat{K} + 2K^{k_i}\}}{\alpha^{k_i}} \delta_{k_i}^2.
\end{aligned} \quad (26)$$

Using (16) (see Lemma 4) and (25) leads to

$$p^{k_i}(z_\ell) - f(x^*) \leq M_{k_i} + \{\hat{K} + 2K^{k_i}\} \delta_{k_i}^2. \quad (27)$$

Since the regression is strict,

$$y(z_\ell, N_{z_\ell}^{k_i}) - p^{k_i}(z_\ell) \leq \beta \sigma(z_\ell, N_{z_\ell}^{k_i}). \quad (28)$$

Using (26), (27), and (28) leads to

$$y(z_\ell, N_{z_\ell}^{k_i}) - f(x^*) \leq M_{k_i} + \{\hat{K} + 2K^{k_i}\} \delta_{k_i}^2 + \beta \left[\frac{\{\hat{K} + 2K^{k_i}\}}{\alpha^{k_i}} \delta_{k_i}^2 \right]. \quad (29)$$

In the second case, $s_c^{k_i}(z) > s_d^{k_i}(x_j)$, then by the construction of M_{k_i} [see (16)],

$$s_d^{k_i}(x_j) - f(x^*) = M_{k_i}. \quad (30)$$

Moreover, since iteration k_i is mesh refinement, then the sampling $N_j \geq \gamma 2^\ell$. Thus, using Assumption 3,

$$\sigma(x_j, N_{x_j}^{k_i}) \leq \alpha \gamma^{-\theta} 2^{-\theta \ell}. \quad (31)$$

Furthermore, the regression $p^{k_i}(x)$ is strict which leads to:

$$y(x_j, N_{x_j}^{k_i}) - s_d^{k_i}(x_j) \leq (\beta + \alpha^{k_i}) \sigma(x_j, N_{x_j}^{k_i}) \quad (32)$$

Using (30), (31), and (32) leads to

$$y(x_j, N_{x_j}^{k_i}) - f(x^*) \leq M_{k_i} + (\beta + \alpha^{k_i}) \alpha \gamma^{-\theta} 2^{-\theta \ell}. \quad (33)$$

Note that η^{k_i} is the candidate point at iteration k_i . Thus, using (26), (29), (31), and (33), and the construction of candidate point (see Definition 5),

$$\begin{aligned} & y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) - f(x^*) + \alpha^{k_i} \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \leq M_{k_i} \\ & + \max \left\{ (\beta + \alpha^{k_i}) \alpha \gamma^{-\theta} 2^{-\theta \ell}, (\hat{K} + 2K^{k_i}) \delta_{k_i}^2 + \beta \left[\frac{(\hat{K} + 2K^{k_i})}{\alpha^{k_i}} \delta_{k_i}^2 \right] \right\} \\ & + \alpha^{k_i} \max \left\{ \alpha \gamma^{-\theta} 2^{-\theta \ell}, \frac{(\hat{K} + 2K^{k_i})}{\alpha^{k_i}} \delta_{k_i}^2 \right\}. \end{aligned} \quad (34)$$

On the other hand,

$$\delta_{k_i} = \frac{\|b - a\|}{2^{\ell_0 + i}}, \quad \alpha^{k_i} = \alpha^0 + i \alpha^\delta, \quad K^{k_i} = K_0 2^i, \quad \ell^{k_i} = \ell^0 + i. \quad (35)$$

By substituting (35) in (34) and using (17) (see Lemma 4), (23a) is verified. Furthermore, using Assumption 2, we have

$$|y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) - f(\eta^{k_i})| \leq E(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \leq Q. \quad (36)$$

Thus, using (23a), $f(\eta^{k_i}) - f(x^*) > 0$, and (36) leads to

$$\limsup_{i \rightarrow \infty} \left\{ -Q + \alpha^{k_i} \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \right\} \leq 0.$$

Since $\sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \geq 0$ and $\lim_{i \rightarrow \infty} \alpha^{k_i} = \infty$, (23b) is verified. $\square$

Theorem 1 Consider η^k as the candidate point at iteration k of Algorithm 2, then

$$\lim_{k \rightarrow \infty} f(\eta^k) = f(x^*), \quad (37)$$

where x^* is a global minimizer of $f(x)$.

Proof At any iteration $k > k_1$, take $k_i < k$ as the most recent mesh refinement iteration of Algorithm 2. Then $\eta^{k_i} \in S^k$, and

$$y(\eta^k, N_{\eta^k}^k) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) \leq y(\eta^{k_i}, N_{\eta^{k_i}}^k) + \alpha^k \sigma(\eta^{k_i}, N_{\eta^{k_i}}^k). \quad (38)$$

Using Assumption 2 leads to:

$$\begin{aligned}
|y(\eta^{k_i}, N_{\eta^{k_i}}^k) - y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i})| &\leq E(\sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i})) + E(\sigma(\eta^{k_i}, N_{\eta^{k_i}}^k)), \\
y(\eta^k, N_{\eta^k}^k) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) &\leq y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) + \alpha^k \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \\
&\quad + E(\sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i})) + E(\sigma(\eta^{k_i}, N_{\eta^{k_i}}^k)).
\end{aligned}$$

By construction, since the sampling at η^{k_i} at iteration k is greater than or equal to its sampling at iteration k_i , $\sigma(\eta^{k_i}, N_{\eta^{k_i}}^k) \leq \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i})$. Since the function $E(x)$ is nondecreasing,

$$y(\eta^k, N_{\eta^k}^k) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) \leq y(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) + \alpha^k \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) + 2E(\sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i})).$$

Using (23) in Lemma 5, Assumption 2, and $\alpha^k = \alpha^{k_i} + \alpha^\delta$, leads to:

$$\limsup_{k \rightarrow \infty} y(\eta^k, N_{\eta^k}^k) - f(x^*) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) \leq \limsup_{k \rightarrow \infty} \left\{ \alpha^\delta \sigma(\eta^{k_i}, N_{\eta^{k_i}}^{k_i}) \right\} = 0. \quad (39)$$

Similar to the proof of (23b), it is thus again easy to show

$$\lim_{i \rightarrow \infty} \sigma(\eta^k, N_{\eta^k}^k) = 0.$$

On the other hand, based on Assumption 2, and optimality of $f(x^*)$

$$\begin{aligned}
f(\eta^k) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) - f(x^*) - E(\sigma(\eta^k, N_{\eta^k}^k)) &\leq y(\eta^k, N_{\eta^k}^k) + \alpha^k \sigma(\eta^k, N_{\eta^k}^k) - f(x^*), \\
\lim_{k \rightarrow \infty} E(\sigma(\eta^k, N_{\eta^k}^k)) &= 0, \\
\limsup_{k \rightarrow \infty} f(\eta^k) - f(x^*) &\leq 0.
\end{aligned}$$

Since $f(\eta^k) - f(x^*) \geq 0$, (37) is verified. $\square$

5 Results

We now illustrate the performance of Algorithm 2 on some representative examples. The function $g(x, k)$ in (1a) is assumed to be a discrete-time statistically stationary random ergodic process. In this section, we further assume that $g(x, k)$ is IID in the index k , and that the variation of $g(x, k)$ from the truth function $f(x)$ is homogeneous in x . In particular,

$$g(x, k) = f(x) + v_k \quad \text{where } v_k \sim \mathcal{N}(0, 0.3).$$

In this section, two different test functions for $f(x)$ are considered within the simple feasible domain $L = \{x | 0 \leq x_i \leq 1 \forall i\}$, the *shifted parabolic function*

$$f(x) = \frac{5}{n} \sum_{i=1}^n (x_i - 0.3)^2, \quad (40)$$

with a global minimizer in L of $x_i^* = 0.3$ and a corresponding global minimum of $f(x^*) = 0$, and the *scaled Schwefel fuction*

$$f(x) = 0.83797 - \frac{1}{n} \sum_{i=1}^n x_i \sin(500 |x_i|), \quad (41)$$

with a global minimizer in L of $x_i^* = 0.8419$ and a corresponding global minimum of $f(x^*) = 0$. We will consider these two functions in $n = 1, 2$, and 3 dimensions.

One-dimensional representations of these functions are illustrated in Fig 2: for the shifted parabolic function (40), the truth function (unknown to the optimization algorithm) is a simple parabola, whereas for the scaled Schwefel fuction (41), the truth function is a smooth nonconvex function with four local minima. Note that the perturbations present in several measurements of these functions, computed with finite N_i , result in a complicated, nonsmooth, nonconvex behavior. This paper shows how to efficiently minimize such functions based only on such noisy measurements, automatically refining the measurements (by increasing the sampling) as convergence is approached.

The optimizations are initialized with measurements of sample length $N^0 = 1$ at the vertices of L . Figure 3 illustrates the application of Algorithm 2 after $k = 200$ iterations in the 1D case, taking $N^0 = N^\delta = 1$ additional sample (at either a new measurement point, or at an existing measurement point) at each iteration of the algorithm. In Fig. 3a, the sampling N_i after $k = 100$ iterations (plus the 2 initial sample points, for a total of 202 samples) at the $M = 5$ measured points y_i indicated, enumerated from left to right, is $\{25, 94, 58, 24, 1\}$; in Fig. 3b, the sampling N_i after 202 iterations at the 7 measured points indicated is $\{7, 6, 11, 4, 55, 115, 4\}$. Both results clearly show that the algorithm focuses the bulk of its sampling in the immediate vicinity of the minimum, where the accuracy of the measurements is especially important, while avoiding unnecessary sampling far from the minimum, where the

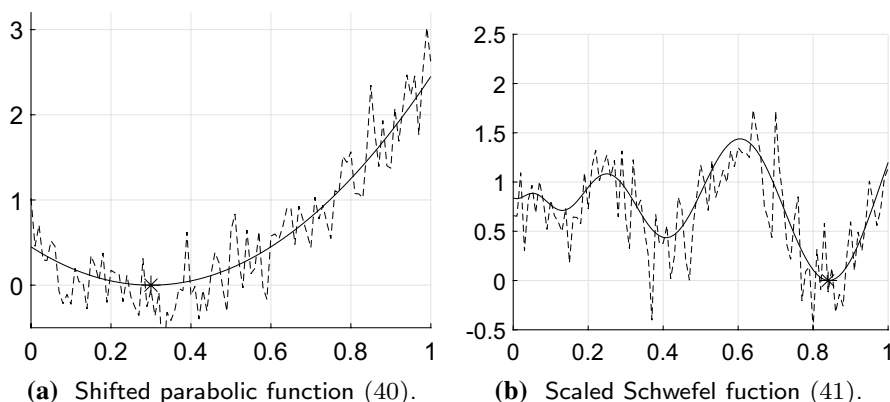


Fig. 2 Illustration of test problems (40) and (41). (solid line) truth function $f(x)$, and (dashed line) a set of measurements y_i computed with a single sample at each measurement, $N_i = 1$. Global minimum of the truth functions are shown by asterisks

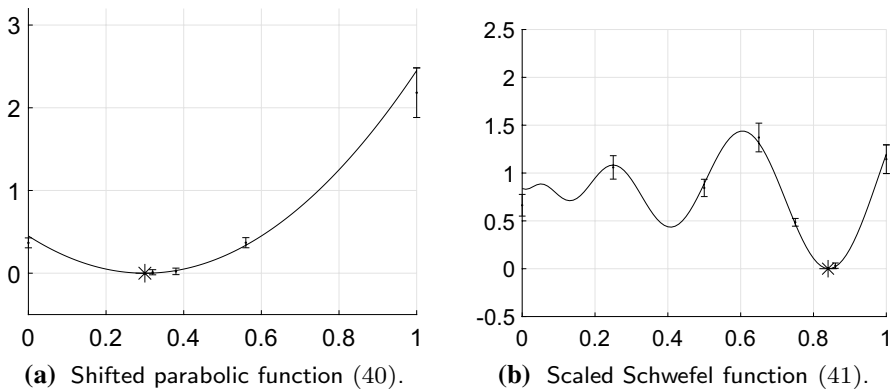


Fig. 3 Illustration of Algorithm 2 on model problems (40) and (41) in 1D after 200 iterations, taking $N^0 = N^\delta = 1$: (solid line) the truth function $f(x)$, and (error bars) the 66% confidence intervals of the measurements. Global Minimum of the truth functions are shown by asterisks

accuracy of the measurements is of reduced importance. It is also seen that more exploration is performed for the scaled Schwefel function than for the shifted parabolic function, as a result of its more complex underlying trend.

Since the function evaluation process in these tests has a stochastic component, Algorithm 2 was next applied an ensemble of twenty separate tests, for both model problems discussed above, in each of three different cases with increasingly higher dimension (that is, $n = 1$, $n = 2$, and $n = 3$). The convergence histories of these simulations are illustrated in Figs. 4 and 5.

To better quantify the performance of the algorithm proposed, we now introduce the following concept.

Definition 6 Assume that the stationary process $g(x, k)$ is IID, and that the nominal variance $\sigma(x_i, 1) = \sigma_0$ for all points $x_i \in L$. As mentioned in Remark 3, denoting N_i as the total number of samples taken at point x_i , the 99.6 confidence intervals of the corresponding measurement y_i given by (5) is $3\sigma_i = 3\sigma(x_i, N_i) = 3\sigma_0/\sqrt{N_i}$. If we assume that all of sampling of the algorithm is performed at a single point, the uncertainty of this single measurement after k samples would thus be $\sigma_0/\sqrt{k}$, which we refer to as the *reference error*. This function is indicated in Figs. 4 and 5 by a solid line of slope $-1/2$ in log-log coordinates.

It is observed that, (see Figs. 4 and 5), in which we have again taken 1 new sample at each iteration, the averaged value of the regret function of Algorithm 2 is eventually diminished to a value close to the reference error. That is, the value of the regret at the end of these optimizations is actually proportional to the uncertainty of a single measurement, assuming that all of the sampling is done at a single point.

Figures 4 and 5 also report the number of datapoints which are considered by the optimization algorithm as the iterations proceed. This number is important in optimization problems for which the function evaluations are obtained from simulations which have an (expensive) initial transient, which must be set

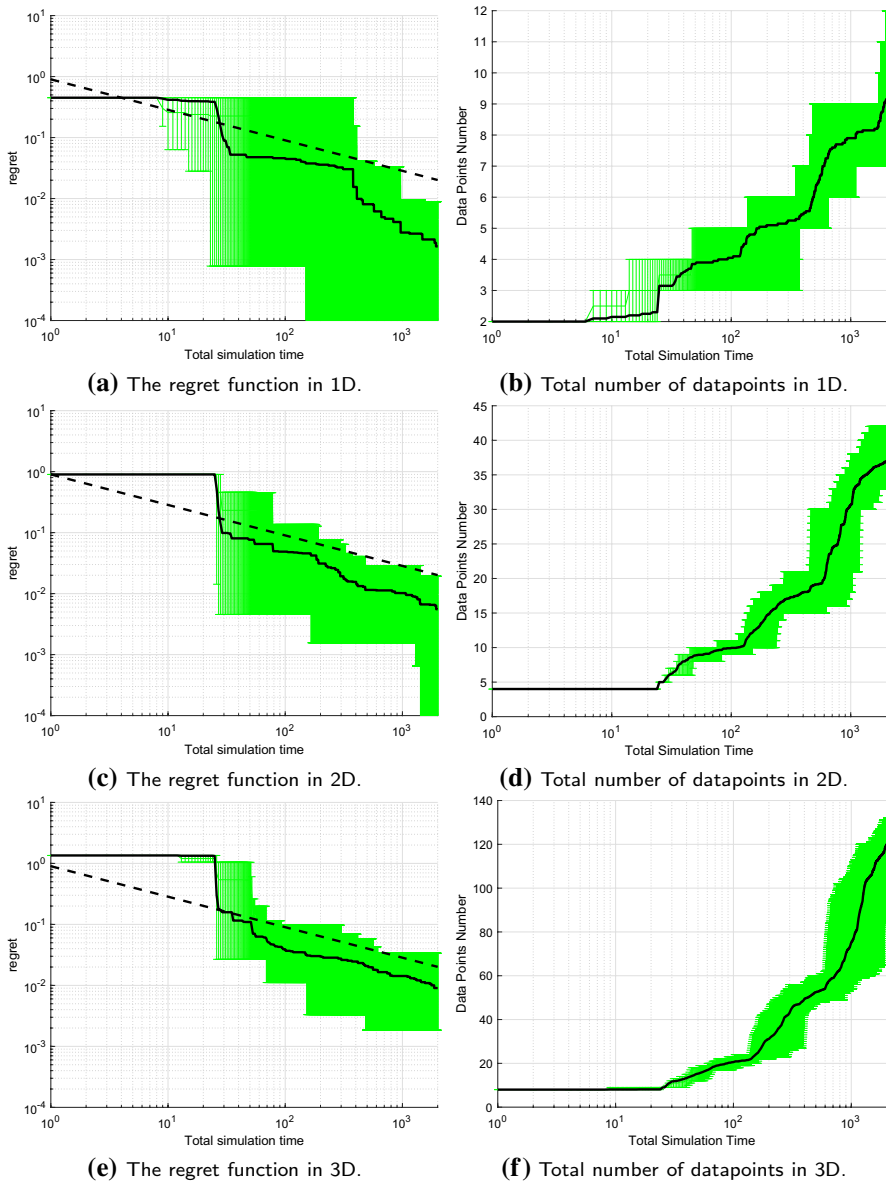


Fig. 4 Implementation of Algorithm 2 on the stochastically-obscured parabolic test problem (40), for twenty different runs. The left figure shows the mean, min, and max value of the regression function over the ensembles. Also plotted at left (dashed bold) is the reference error

aside before sampling the statistic of interest, as discussed further in Remark 2. It is observed, as in the 1D case illustrated in Fig. 3, that the number of data-points that are considered for the shifted parabolic function is less than that for the scaled Schwefel function. Further, the regret function converges faster to the

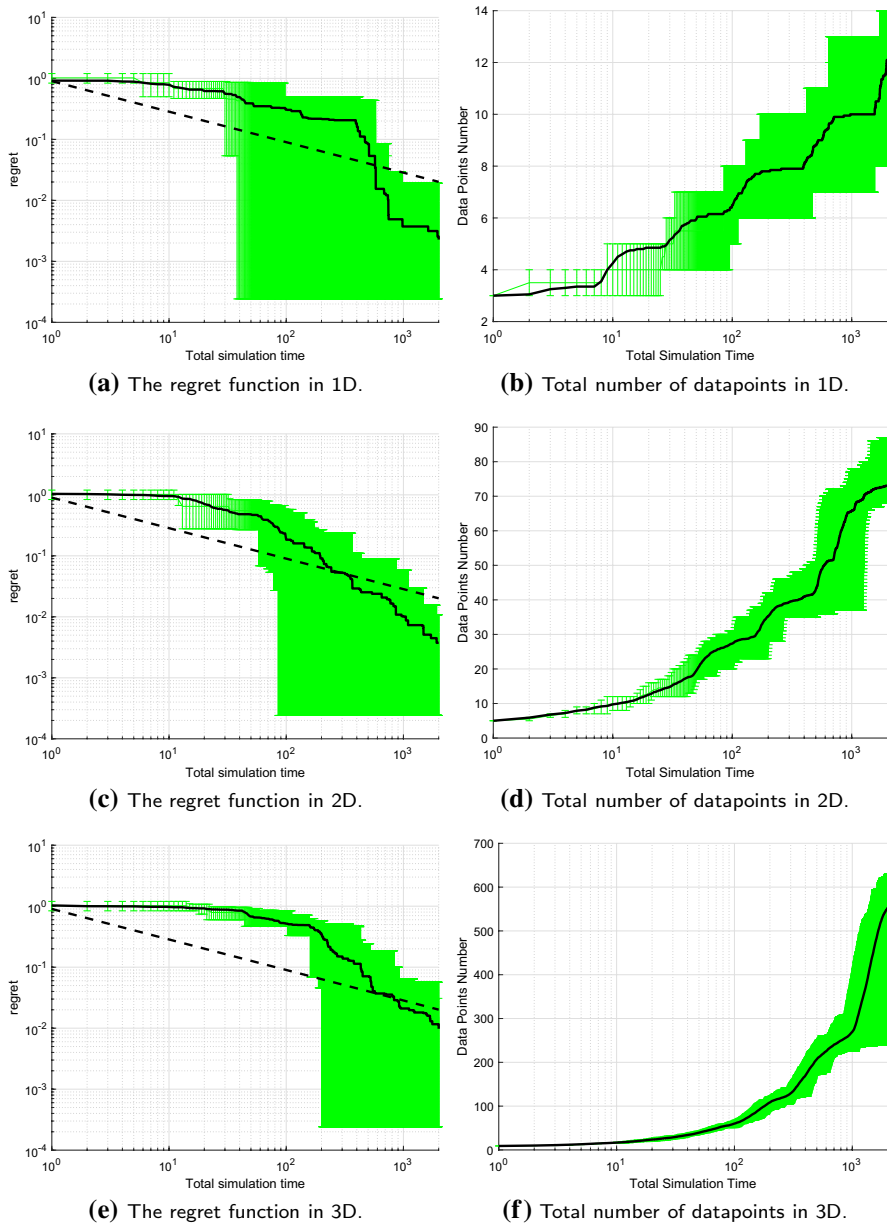


Fig. 5 Implementation of Algorithm 2 on the stochastically-obscured Schwefel test problem (41), for twenty different runs. The left figure shows the mean, min, and max value of the regression function over the ensembles. Also plotted at left (dashed bold) is the reference error

general proximity of the global solution for the shifted parabolic function. This result is reasonable, since the underlying function in the shifted parabolic case is much simpler.

The simulations whose results are shown in Figs. 4 and 5 indicate that the presented algorithm has improved the performance of the optimization process by reducing the simulation time at the datapoints, which are far from the solution. However, the necessity of refining steps needs further considerations. For this purpose, we have applied algorithm 2 with and without refining step on to problem (40) whose results are shown in Fig. 6. For the optimization algorithm without refining steps, the candidate point at each iteration is the minimizer of the regression function in the datapoints. It is observed that by removing the refining steps from the optimization process, convergence speed varies significantly among the ensemble runs. In other words, the presence of the refining step in the optimization process is necessary to ensure consistent convergence performance among the ensemble runs. Another advantage of the refining step is to reduce the total number of datapoints in the optimization process as shown in Fig. 6. This is a significant advantage, especially for those range of optimization problems whose function evaluations are obtained from simulations with a long initial transient time.

6 Application of α -DOGS to estimate the parameters of a Lorenz system

In this section, Algorithm 1 (Δ -DOGS), Algorithm 2 (α -DOGS), and the Surrogate Management Framework (SMF) developed in [15] and implemented in [25] are applied to a representative optimization problem, based on infinite-time-averaged statistics, of the type considered in this paper [see (1)].

The specific problem considered here is derived from the well-known 3-state Lorenz model [21, 41], the dynamics of which exhibit a familiar chaotic behavior that roughly characterizes the bulk flow of a fluid within a hollow torus that is heated from below and cooled from above [10], and is governed by

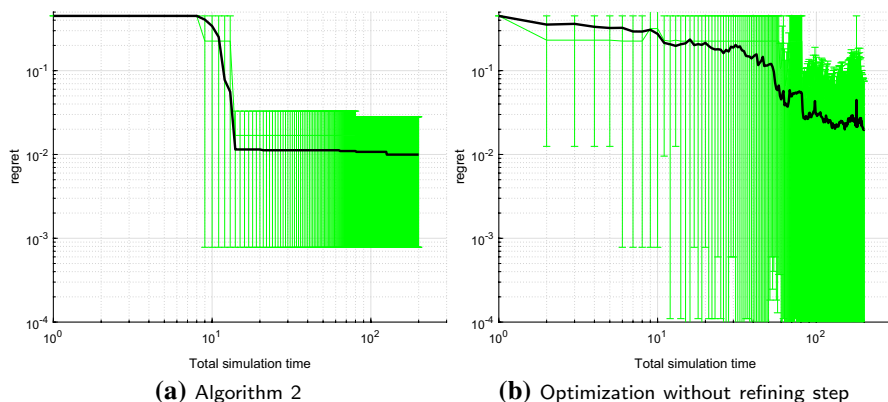


Fig. 6 Implementation of Algorithm 2 with and without refining step on to the optimization problem (40) for twenty different runs. Left figure shows the convergence history of regret function for algorithm 2 with refining step, and right figure shows the results without refining step

$$\frac{d}{dt}X = s(Y - X), \quad (42a)$$

$$\frac{d}{dt}Y = -XZ + \rho X - Y, \quad (42b)$$

$$\frac{d}{dt}Z = XY - \beta Z, \quad (42c)$$

where (X, Y, Z) are the components of the state, which generally moves along a chaotic attractor, and (ρ, β, s) are the three (constant) adjustable parameters which affect various characteristics of this attractor. The infinite-time-averaged mean and standard deviation of the Z component of this ergodic system are given by:

$$\bar{Z} = \lim_{T \rightarrow \infty} \frac{1}{T} \int_{t=0}^T (Z(t)) dt, \quad (43a)$$

$$\hat{Z} = \lim_{T \rightarrow \infty} \sqrt{\frac{1}{T} \int_{t=0}^T (Z(t) - \bar{Z})^2 dt}. \quad (43b)$$

In the optimization problem considered in this section, the value of $s = 10$ is taken as known, and the parameters ρ and β are considered as optimization variables. Using the method developed in this paper (which is based on successive finite-time simulations), we will seek the values of ρ and β which reproduce known values of $\bar{Z}$ and $\hat{Z}$ in the infinite time averaged statistics of the Lorenz system (42). Towards this end, the cost function considered in this section is

$$f(x) = |(\bar{Z} - 23.57)| + |(\hat{Z} - 8.67)|, \quad x = (\rho, \beta). \quad (44)$$

Note that the value of $\hat{Z}$ and $\bar{Z}$ are functions of $\{\rho, \beta\}$. The search domain for $\{\rho, \beta\}$ is taken as $24 \leq \rho \leq 29.15$ and $1.8 \leq \beta \leq 4$; note that the Lorenz system (42) exhibits a statistically stationary ergodic behavior everywhere within this search domain [21]. The optimal solution to this optimization problem is known to be approximately $\rho = 28$, $\beta = 2.667$.

The cost function given in (44) might initially appear to be in a slightly different form than that given in (1a). However, it has the same essential structure, in that $f(x)$ can be approximated with increasing accuracy by increasing the sampling (as well as the computational cost) of any given measurement. As a result, Algorithm (2) can be applied to this problem directly, given a sufficiently representative uncertainty quantification (UQ) procedure for the finite-time-averaged approximations of the infinite-time-averaged statistics of interest. In this work, for the purpose of illustration, we will use the simple UQ approach proposed in Appendix 2, which proves to be adequate for our purposes here; improved UQ approaches developed and discussed elsewhere could certainly be used instead.

To numerically simulate the ODE given in (42), we use a standard RK4 method with a uniform timestep of $h = 0.05$; this approach provides a reasonably

small time discretization error [21, 29] for this problem. [Using the same timestep h in all simulations during the optimization process is a drawback of the α -DOGS optimization algorithm as developed in this paper; this limitation will be addressed in a future paper, which is currently under development.]

Initial conditions near the attractor are taken for all simulations, and (for simplicity) the first 2600 timesteps (up to $T = 13$) are deleted from all simulations, in order to begin time averaging after the system has closely approached the attractor itself. [The interesting problem of automating the detection of such “initial transients”, during which the chaotic system approaches the attractor, is also deferred to a future paper.]

Note that all optimizations performed in this section are terminated when

$$|\hat{f}(x, T) - f(x^*)| \leq 0.04 \quad \text{and} \quad (45)$$

$$\sigma(x, T) \leq 0.02, \quad (46)$$

where $\hat{f}(x, T)$ and $\sigma(x, T)$ are the estimates and uncertainty at point x , and $f(x^*) = 0$ is the global minimum of $f(x)$.

According to the procedure developed in Appendix 2, $T = 2513$ (502600 timesteps) is the minimum simulation time required to achieve the target uncertainty in (46). As a result, all simulations of Δ -DOGS and SMF use fixed time-averaging lengths of $T = 2513$. In contrast, the time averaging length used by α -DOGS for each datapoint computed during the optimization process is controlled by the optimizer itself, and depends on two parameters:

- (a) The averaging length during identifying sampling iterations, which is denoted by $T_0 = 20$ (4000 timesteps). Note that the first 2513 timesteps are not included in time-averaging process, as they are in the startup transient.
- (b) The additional averaging length during supplemental sampling iterations, which is denoted by $T_1 = 7$ (1400 timesteps).

The results of applying the three optimization algorithms considered to problem (44) may summarized as follows:

- (a) α -DOGS requires greatly reduced (up to two orders of magnitude) total averaging time as compared with the other two algorithms considered (see Fig. 7).
- (b) Δ -DOGS uses fewer actual datapoints than α -DOGS (see Fig. 8) to achieve a desired degree of convergence; however, since the time-averaging length at all datapoint is much higher in Δ -DOGS, the actual rate of convergence, in terms of computation time, is greatly improved using α -DOGS.
- (c) For α -DOGS, the required averaging times at the datapoints which have reduced cost function values are significantly greater than the required averaging times at the datapoints that are far from the desired solution (see Fig. 9). As a result, the computational cost of the global exploration process during the optimization is significantly reduced using α -DOGS.

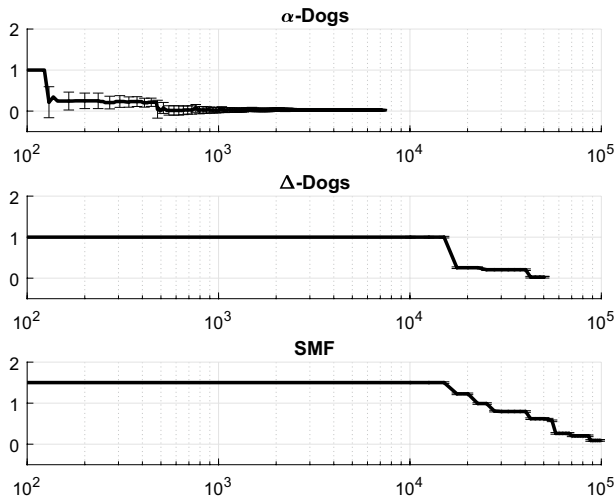


Fig. 7 Best measurement versus total averaging length. Convergence history of optimization algorithm for optimization problem (44) based on Lorenz system

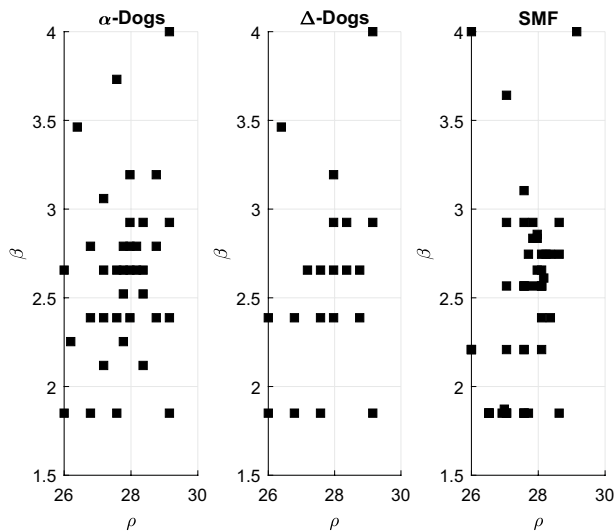


Fig. 8 Location of the datapoints considered during the optimization of (44) based on the Lorenz system

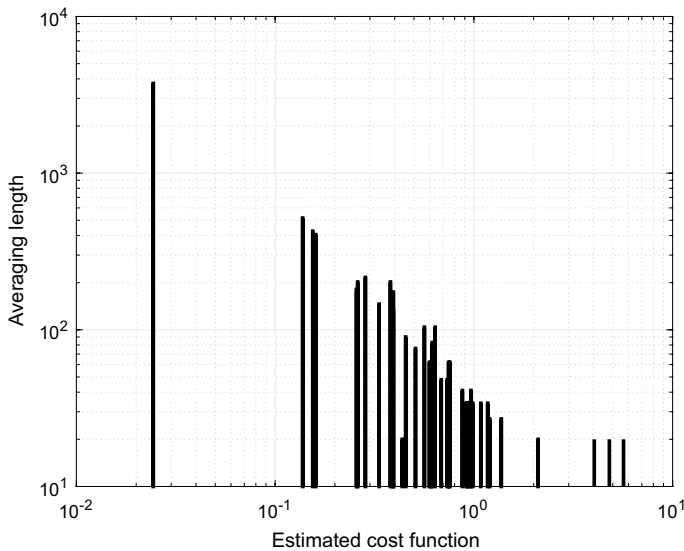


Fig. 9 Relationship between the averaging length and cost function value for α -DOGS applied on optimization problem (44)

7 Conclusions

This paper presents a new optimization algorithm, dubbed α -DOGS, for the minimization of functions given by the infinite-time average of stationary ergodic processes in the computational or experimental setting. Two search functions are considered at each iteration. The first is a continuous search functions, $s_c^k(x)$, defined over the entire feasible space $x \in L$, combining a strict regression $p^k(x)$ of the available data-points together with a remoteness function characterizing the distance of any given point in the feasible domain from the nearest measurements, and built on the framework of a Delaunay triangulation over all available measurements at that iteration. The second is a discrete search function, $s_d^k(x_i)$, defined over the available measurements $x_i \in S^k$. A comparison between the minima of these two search functions is made in order to decide between further sampling (and, therefore, refining) an existing measurement, or sampling at a new point in parameter space. The method developed builds closely on the Delaunay-based Derivative-free Optimization via Global Surrogates algorithm, dubbed Δ -DOGS, proposed in [12–14]. Convergence of the algorithm is established in problems for which

- The underlying truth (infinite-time averaged) function, as well as the regressions computed at each iteration k , are twice differentiable.
- The stationary process $g(x, k)$ upon which the truth function $f(x)$ is generated, in (1a), is ergodic, and the convergence of the averaging process to the underlying

truth function is bounded by a monotonic function of a computable uncertainty function (see Assumption 2).

- (c) The uncertainty of the time averaging process decays exponentially to zero (see Assumption 3); this is true for almost all stationary models of random processes.

The α -DOGS algorithm performs and refines measurements with different amounts of sampling in different locations in the feasible region of parameter space as necessary. By doing so, the total cost of the optimization process is substantially reduced as compared with using existing derivative-free optimization strategies, with the same amount of sampling at different locations in parameter space. Computational experiments demonstrate that the algorithm developed ultimately devotes most of its sampling time to points in parameter space near to the global minimum. Further, these computational experiments indicate that the regret function (see Definition 5) eventually diminishes to a value that is actually substantially less than the uncertainty of a single measurement, assuming that all of the sampling is done at a single point.

In future work, the α -DOGS algorithm will be applied to additional benchmark and application-based optimization problems, including shape optimization for airfoils and hydrofoils. For problems in which the function is determined computationally (from, e.g., numerical simulations of turbulent flows), the extension of the present framework to, as convergence is approached, simultaneously (a) refine the computational grid, and (b) increase the measurement sampling, is also under development.

Acknowledgements We would like to thank reviewers for their insightful comments on the paper, as these comments led us to an improvement of the work. Moreover, We gratefully acknowledge Prof. Philip Gill and Prof. Alison Marsden for their collaborations and funding from AFOSR FA 9550-12-1-0046, Cymer Center for Control Systems & Dynamics, and Leidos corporation in support of this work.

Appendix 1: Polyharmonic spline regression

The algorithm described in this paper depends upon a smooth regression $p^k(x)$ (see Assumption 1). The best technique for computing the regression is problem dependent. As with [12–14], a key advantage of our Delaunay-based approach in the present work is that it facilitates the use of *any* suitable regression technique, subject to it satisfying the “strict” regression property given in Definition 4. Since our numerical tests all implement the polyharmonic spline regression technique, the derivation of this regression technique is briefly explained in this appendix; additional details may be found in [46].

The polyharmonic spline regression $p(x)$ of a function $f(x)$ in $\mathbb{R}^n$ is defined as a weighted sum of a set of radial basis functions $\varphi(r)$ built around the location of each measurement point, plus a linear function of x :

$$p(x) = \sum_{i=1}^N w_i \varphi(r) + v^T \begin{bmatrix} 1 \\ x \end{bmatrix}, \quad (47)$$

where $\varphi(r) = r^3$ and $r = \|x - x_i\|$.

The weights w_i and v_i represent N and $n+1$ unknowns. Assume that $\{y(x_1), y(x_2), \dots, y(x_n)\}$ is the set of measurements, with standard deviations $\{\sigma_1, \sigma_2, \dots, \sigma_n\}$. The w_i and v_i coefficients are computed by minimizing the following objective function, which expresses a tradeoff between the fit to the observed data and the smoothness of the regressor:

$$L_p(x) = \sum_{i=1}^N \left[\frac{(p(x_i) - y(x_i))^2}{\sigma_i^2} \right] + \lambda \int_B |\nabla^m p(x)|, \quad (48)$$

where B is a large box domain containing all of the x_i , and $\nabla^m p(x)$ is the vector including all m derivatives of $p(x)$ (see [20]). It is shown in [20] that the first-order optimality condition for the objective function (48) is as follows:

$$p(x_i) - y(x_i) + \rho \sigma_i^2 w_i = 0, \quad \forall 1 \leq i \leq N, \quad (49)$$

where ρ is a parameter proportional to λ . In summary, the coefficient of the regression can be derived by solving:

$$\begin{bmatrix} F & V^T \\ V & 0 \end{bmatrix} \begin{bmatrix} w \\ v \end{bmatrix} = \begin{bmatrix} f(x_i) \\ 0 \end{bmatrix}, \quad (50)$$

$$F_{ij} = \varphi(\|x_i - x_j\|) + \rho \delta_{ij} \sigma_i^2, \quad V = \begin{bmatrix} 1 & 1 & \dots & 1 \\ x_1 & x_2 & \dots & x_N \end{bmatrix},$$

where δ_{ij} is the Kronecker delta.

The problem which is left to solve when computing the regression is to find an appropriate value of $\rho \in [0, \infty)$. Solving (50) for any value of ρ gives a unique regression, denoted $p(x, \rho)$. The parameter ρ is then obtained by a predictive mean-square error criteria developed in §4.4 in [46], which is given by imposing the following condition:

$$T(\rho) = \sum_{i=1}^N \left[\frac{p(x_i, \rho) - y(x_i)}{\sigma_i} \right]^2 = 1. \quad (51)$$

For $\rho \rightarrow \infty$, $w_i \rightarrow 0$, and the solution of (50) is a weighted mean-square linear regression, which is obtained by solving (51). If $T(\infty) \leq 1$, we take this linear regression as the best current regression for the available data. Otherwise, we have $T(\infty) > 1$ and (by construction) $T(0) = 0$; thus, (51) has a solution with finite $\rho > 0$, which gives the desired regression.

If $T(\infty) > 1$, we thus seek a ρ for which for $T(\rho) = 1$. Following [46], using (50), (51) simplifies to:

$$T(\rho) = \rho^2 \left(\sum_{i=1}^N w_i \sigma_i \right)^2 = 1, \quad (52)$$

where $w_{i,\rho}$ is the w_i which is obtained by solving (50). Define Dw and Dv as the vectors whose i -th elements are the derivatives of w_i and v_i with respect to ρ , then

$$T'(\rho) = \rho^2 \sum_{i=1}^N w_i D w_i \sigma_i^2 + 2 \rho \left(\sum_{i=1}^N w_{i,\rho} \sigma_i \right)^2,$$

$$\begin{bmatrix} F & V^T \\ V & 0 \end{bmatrix} \begin{bmatrix} Dw \\ Dv \end{bmatrix} + \begin{bmatrix} \rho \Sigma_2 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} w \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix},$$

where Σ_2 is a diagonal matrix whose i -th diagonal element is $\rho \sigma_i^2$. Therefore, the analytic expression for the derivative of $T(\rho)$ is available. Thus, (51) can be solved quickly using Newton's method.

The regression process presented here, imposing (1) as suggested by [46], is designed to obtain a regression which is reasonably smooth. However, there is no guarantee that this particular regression satisfies the strictness property required in the present work (see Definition 4). Note, however, that by imposing $\rho = 0$, the regression is made strict for arbitrary small β . Thus, to satisfy strictness for a given finite β , the value of ρ must sometimes be decreased from that which satisfies (1), as necessary.

Appendix 2: UQ for finite-time-averaging of the Lorenz system

This appendix summarizes briefly the simple empirical approach that is used in this paper to quantify the uncertainty of the cost function (44) when it is estimated using a finite time average.

In the method used, we simply simulated the Lorenz system (42) 30 independent times with various initial values for (X, Y, Z) . The cost function was then approximated using these different simulation lengths, and the standard deviation of the estimations obtained using various simulation lengths was calculated. The simple model given by $A/\sqrt{T}$ for the uncertainty was found to fit this empirical calculation quite well, as shown in Fig. 10.

Appendix 3: Application of α -DOGS on higher dimensional problems

We now consider the challenge of increasing the dimension of the problem considered. This is rather difficult, as the algorithm developed herein (specifically, the computational cost of computing the Delaunay triangulation) scales rather poorly as the dimension of the problem is increased. In this appendix, the convergence behavior on a 5-dimensional problem is thus investigated.

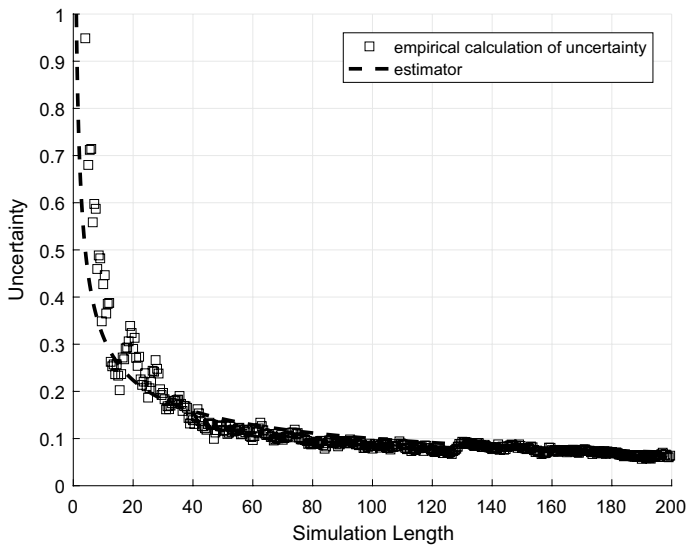


Fig. 10 Uncertainty quantification (UQ) for finite-time-average approximations of the infinite-time-average statistic of interest in the cost function related to the Lorenz model. The uncertainty quantification model of $A/\sqrt{T}$ is found to give a very good empirical fit

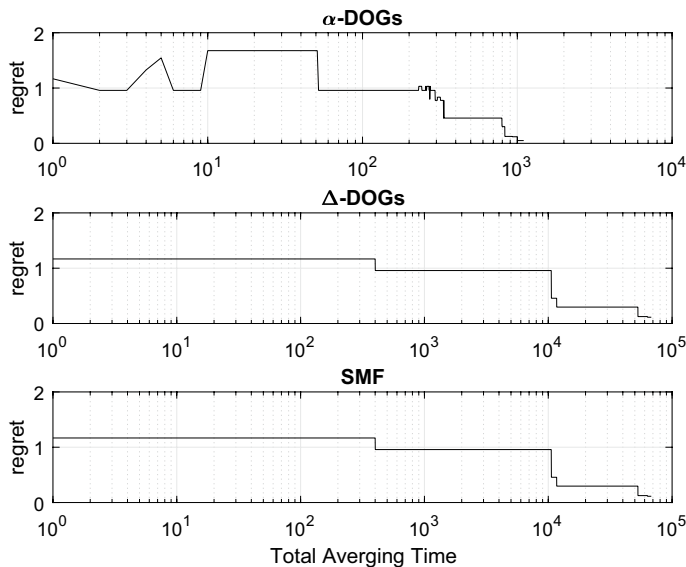


Fig. 11 Convergence history of optimization algorithms for (41) for $n = 5$

The test problem that is considered here is the stochastically-obscured scaled Schwefel function (41) for $n = 5$. As in Sect. 6, performance of the α -DOGS, Δ -DOGS, and SMF algorithms are characterized. In this investigation, for all datapoint

in Δ -DOGs and SMF algorithm, a fixed averaging time of $T = 100$ was used. Figure 11 shows the convergence history of the regret function for each method. It is observed that the required total averaging time for the optimization problem with α -DOGS is significantly better than the other methods considered.

To scale to even higher dimensional problems, the practical (though, perhaps, not provably globally convergent) idea of using only *approximate* Delaunay triangulations has been proposed, and will be considered in future work.

References

1. Alimo, S., Beyhaghi, P., Meneghello, G., Bewley, T.: Delaunay-based optimization in CFD leveraging multivariate adaptive polyharmonic splines (maps). In: 58th AIAA/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference, p. 0129 (2017)
2. Alimo, S.R., Beyhaghi, P., Bewley, T.R.: Optimization combining derivative-free global exploration with derivative-based local refinement. In: 2017 IEEE 56th Annual Conference on Decision and Control (CDC), pp. 2531–2538. IEEE (2017)
3. Amaioua, N., Audet, C., Conn, A.R., Le Digabel, S.: Efficient solution of quadratically constrained quadratic subproblems within the mesh adaptive direct search algorithm. *Eur. J. Oper. Res.* **268**(1), 13–24 (2018)
4. Audet, C., Conn, A.R., Le Digabel, S., Peyrega, M.: A progressive barrier derivative-free trust-region algorithm for constrained optimization. *Comput. Optim. Appl.* **71**(2), 307–329 (2018)
5. Audet, C., Hare, W.: *Derivative-Free and Blackbox Optimization*. Springer, Berlin (2017)
6. Audet, C., Tribes, C.: Mesh-based Nelder–Mead algorithm for inequality constrained optimization. *Comput. Optim. Appl.* **71**(2), 331–352 (2018)
7. Awad, H.P., Glynn, P.W.: On an initial transient deletion rule with rigorous theoretical support. In: *Proceedings of the 38th Conference on Winter Simulation*, pp. 186–191. Winter Simulation Conference (2006)
8. Beran, J.: *Statistics for Long-Memory Processes*, vol. 61. CRC Press, Boca Raton (1994)
9. Beran, J.: Maximum likelihood estimation of the differencing parameter for invertible short and long memory autoregressive integrated moving average models. *J. R. Stat. Soc. Ser. B (Methodol.)* **57**, 659–672 (1995)
10. Bewley, T.R., Moin, P., Temam, R.: Dns-based predictive control of turbulence: an optimal benchmark for feedback algorithms. *J. Fluid Mech.* **447**, 179–225 (2001)
11. Beyhaghi, P., Alimohammadi, S., Bewley, T.: Uncertainty quantification of the time averaging of a statistics computed from numerical simulation of turbulent flow (2018). arXiv preprint [arXiv:1802.01056](https://arxiv.org/abs/1802.01056)
12. Beyhaghi, P., Bewley, T.: Implementation of Cartesian grids to accelerate Delaunay-based derivative-free optimization. *J. Glob. Optim.* **69**(4), 927–949 (2017)
13. Beyhaghi, P., Bewley, T.R.: Delaunay-based derivative-free optimization via global surrogates, part II: convex constraints. *J. Glob. Optim.* **66**, 383–415 (2016)
14. Beyhaghi, P., Cavaglieri, D., Bewley, T.: Delaunay-based derivative-free optimization via global surrogates, part I: linear constraints. *J. Glob. Optim.* **66**, 1–52 (2015)
15. Booker, A.J., Dennis Jr., J.E., Frank, P.D., Serafini, D.B., Torczon, V., Trosset, M.W.: A rigorous framework for optimization of expensive functions by surrogates. *Struct. Optim.* **17**(1), 1–13 (1999)
16. Bubeck, S., Munos, R., Stoltz, G., Szepesvari, C.: X-armed bandits. *J. Mach. Learn. Res.* **12**, 1655–1695 (2011)
17. Conn, A.R., Scheinberg, K., Vicente, L.N.: Global convergence of general derivative-free trust-region algorithms to first- and second-order critical points. *SIAM J. Optim.* **20**(1), 387–415 (2009)
18. Conn, A.R., Scheinberg, K., Vicente, L.N.: *Introduction to Derivative-Free Optimization*, vol. 8. SIAM, Philadelphia (2009)

19. Deng G., Ferris, M.C.: Extension of the direct optimization algorithm for noisy functions. In: Proceedings of the 39th Conference on Winter Simulation: 40 Years! The Best is Yet to Come, pp. 497–504. IEEE Press (2007)
20. Duchon, J.: Splines minimizing rotation-invariant semi-norms in Sobolev spaces. In: Constructive Theory of Functions of Several Variables, pp. 85–100. Springer (1977)
21. Goluskin, D.: Bounding averages rigorously using semidefinite programming: mean moments of the Lorenz system. *J. Nonlinear Sci.* **28**, 621–651 (2017)
22. Jones, D.R., Perttunen, C.D., Stuckman, B.E.: Lipschitzian optimization without the Lipschitz constant. *J. Optim. Theory Appl.* **79**(1), 157–181 (1993)
23. Kleinberg, R., Slivkins, A., Upfal, E.: Multi-armed bandits in metric spaces. In: Proceedings of the Fortieth Annual ACM Symposium on Theory of Computing, pp. 681–690. ACM (2008)
24. Lasserre, J.B.: An Introduction to Polynomial and Semi-algebraic Optimization, vol. 52. Cambridge University Press, Cambridge (2015)
25. Marsden, A.L., Wang, M., Dennis Jr., J.E., Moin, P.: Optimal aeroacoustic shape design using the surrogate management framework. *Optim. Eng.* **5**(2), 235–262 (2004)
26. Marsden, A.L., Wang, M., Dennis Jr., J.E., Moin, P.: Suppression of vortex-shedding noise via derivative-free shape optimization. *Phys. Fluids* **16**(10), 83–86 (2004)
27. Nie, J., Yang, L., Zhong, S.: Stochastic polynomial optimization. In: Optimization Methods and Software, pp. 1–19 (2019)
28. Norkin, V.I., Pflug, G.C., Ruszczyński, A.: A branch and bound method for stochastic global optimization. *Math. Program.* **83**(1–3), 425–450 (1998)
29. Oliver, T.A., Malaya, N., Ulerich, R., Moser, R.D.: Estimating uncertainties in statistics computed from direct numerical simulation. *Phys. Fluids* (1994–present) **26**(3), 035101 (2014)
30. Picheny, V., Ginsbourger, D., Richet, Y., Caplin, G.: Quantile-based optimization of noisy computer experiments with tunable precision. *Technometrics* **55**(1), 2–13 (2013)
31. Quan, N., Yin, J., Ng, S.H., Lee, L.H.: Simulation optimization via kriging: a sequential search using expected improvement with computing budget constraints. *IIE Trans.* **45**(7), 763–780 (2013)
32. Rasmussen, C.E.: Gaussian Processes for Machine Learning. Citeseer (2006)
33. Rullière, D., Faleh, A., Planchet, F., Youssef, W.: Exploring or reducing noise? *Struct. Multidiscip. Optim.* **47**(6), 921–936 (2013)
34. Salesky, S.T., Chamecki, M., Dias, N.L.: Estimating the random error in eddy-covariance based fluxes and other turbulence statistics: the filtering method. *Bound.-Layer Meteorol.* **144**(1), 113–135 (2012)
35. Sankaran, S., Audet, C., Marsden, A.L.: A method for stochastic constrained optimization using derivative-free surrogate pattern search and collocation. *J. Comput. Phys.* **229**(12), 4664–4682 (2010)
36. Schonlau, M., Welch, W.J., Jones, D.R.: A data-analytic approach to Bayesian global optimization. In: Department of Statistics and Actuarial Science and The Institute for Improvement in Quality and Productivity, 1997 ASA Conference (1997)
37. Sethi, S.P., Zhang, Q., Zhang, H.-Q.: Average-Cost Control of Stochastic Manufacturing Systems, vol. 54. Springer, Berlin (2005)
38. Slivkins, A.: Multi-armed bandits on implicit metric spaces. In: Advances in Neural Information Processing Systems, pp. 1602–1610 (2011)
39. Snoek, J., Larochelle, H., Adams, R.P.: Practical Bayesian optimization of machine learning algorithms. In: Advances in Neural Information Processing Systems, pp. 2951–2959 (2012)
40. Srinivas, N., Krause, A., Kakade, S.M., Seeger, M.W.: Information-theoretic regret bounds for Gaussian process optimization in the bandit setting. *IEEE Trans. Inf. Theory* **58**(5), 3250–3265 (2012)
41. Stewart, I.: Mathematics: The Lorenz attractor exists. *Nature* **406**(6799), 948 (2000)
42. Talgorn, B., Le Digabel, S., Kokkolaras, M.: Statistical surrogate formulations for simulation-based design optimization. *J. Mech. Des.* **137**(2), 021405 (2015)
43. Talnikar, C., Blonigan, P., Bodart, J., Wang, Q.: Parallel optimization for les. In: Proceedings of the Summer Program, p. 315 (2014)
44. Theunissen, R., Di Sante, A., Riethmuller, M.L., Van den Braembussche, R.A.: Confidence estimation using dependent circular block bootstrapping: application to the statistical analysis of PIV measurements. *Exp. Fluids* **44**(4), 591–596 (2008)

45. Valko, M., Carpentier, A., Munos, R.: Stochastic simultaneous optimistic optimization. In: International Conference on Machine Learning, pp. 19–27 (2013)
46. Wahba, G.: Spline Models for Observational Data, vol. 59. SIAM, Philadelphia (1990)
47. Zhao, M., Alimo, S.R., Bewley, T.R.: An active subspace method for accelerating convergence in Delaunay-based optimization via dimension reduction. In: 2018 IEEE Conference on Decision and Control (CDC), pp. 2765–2770. IEEE (2018)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.